



(21) 申请号 202311760197.4

(22) 申请日 2023.12.20

(65) 同一申请的已公布的文献号

申请公布号 CN 117709969 A

(43) 申请公布日 2024.03.15

(73) 专利权人 华南理工大学

地址 510640 广东省广州市天河区五山路
381号专利权人 人工智能与数字经济广东省实验
室(广州)

(72) 发明人 张通 邓忠易 陈俊龙

(74) 专利代理机构 广州市华学知识产权代理有
限公司 44245

专利代理师 霍健兰

(51) Int.Cl.

G06Q 30/01 (2023.01)

G06N 3/0455 (2023.01)

G06N 3/0895 (2023.01)

G06N 3/084 (2023.01)

(56) 对比文件

CN 111444721 A, 2020.07.24

CN 114756658 A, 2022.07.15

审查员 史金红

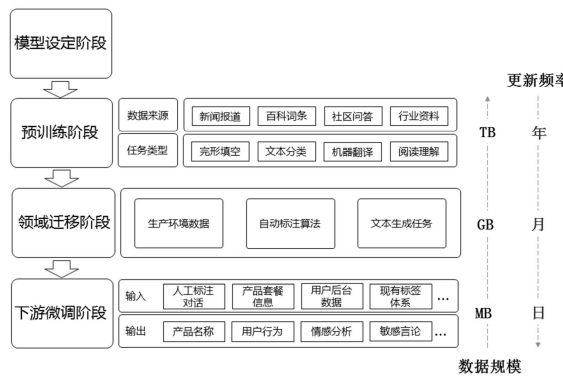
权利要求书2页 说明书8页 附图3页

(54) 发明名称

面向客服场景的生成匹配式大模型构建方法、介质及设备

(57) 摘要

本发明提供了一种面向客服场景的生成匹配式大模型构建方法、介质及设备；该方法包括依次执行的模型设定阶段、预训练阶段、领域迁移阶段和下游微调阶段；预训练阶段是指：采用跨领域中文语料库的文本作为样本，对智慧客服大模型的大模型基座进行预训练；领域迁移阶段是指：采用客服场景数据作为样本；对智慧客服大模型的大模型基座进行弱监督训练；下游微调阶段是指：采用客服场景人工标注数据作为样本，对智慧客服大模型进行训练以学习新业务的相关知识。该方法分阶段逐步地实现和优化大模型的功能，使其具备深度挖掘大规模客服文本数据知识的能力，同时对新增的业务需求和变化的业务内容具备精准迁移和快速扩展的能力。



1. 一种面向客服场景的生成匹配式大模型构建方法,其特征在于:包括依次执行的模型设定阶段、预训练阶段、领域迁移阶段和下游微调阶段;

所述模型设定阶段,是指:设定基于Transformer架构的智慧客服大模型;所述智慧客服大模型包括大模型基座和连接在大模型基座输出端的映射模块;

所述预训练阶段,是指:采用跨领域中文语料库的文本作为样本,对智慧客服大模型的大模型基座进行预训练,以使智慧客服大模型具备强泛化能力;

所述领域迁移阶段,是指:采用客服场景数据作为样本;利用客服场景数据中的精标数据和通过自动标注得到的伪标签数据,对智慧客服大模型的大模型基座进行弱监督训练,以使智慧客服大模型具备强领域知识;

所述下游微调阶段,是指:采用客服场景人工标注数据作为样本;样本根据业务需求拼接对应的业务需求提示模板,得到任务数据;利用任务数据对智慧客服大模型进行训练,以学习新业务的相关知识;

所述领域迁移阶段,包括如下步骤:

X1、采集客服场景数据;将客服场景数据分类为用户信息一、产品信息、客服对话;客服对话包括语音转录的对话文本一;

对话文本一包括有标签的对话文本一和无标签的对话文本一;有标签的对话文本一数量>无标签的对话文本一数量;

X2、当对话文本一为有标签的对话文本一时,以“对话文本-标签”形式形成精标数据;

当对话文本一为无标签的对话文本一时,获取对话文本一对应的用户信息一;根据上下文模板,将用户信息一转换成文字描述,并与对话文本一进行拼接,得到对话拼接文本;将对话拼接文本进行自动标注处理,得到对话文本一对应的摘要;将对话文本一和摘要组合,以“对话文本-摘要”形式形成伪标签数据,以实现数据增强;

X3、将标签数据和伪标签数据分别输入智慧客服大模型的大模型基座;在一次前向传播过程中,计算得到标签数据和伪标签数据的损失梯度,并比较损失梯度的相似程度:

若损失梯度方向一致,则判定为训练有效,之后进行后向传播,以完成弱监督训练;

否则将伪标签数据的损失梯度置零,之后进行后向传播,以完成弱监督训练;

所述下游微调阶段,包括如下分步骤:

Y1、采集客服场景人工标注数据;人工标注数据包括有标签的对话文本二;获取对话文本二对应的用户信息二;根据业务需求设定业务需求提示模板;根据所述上下文模板,将用户信息二转换成文字描述,并与对话文本二进行拼接,之后与业务需求提示模板拼接,得到任务数据;

Y2、设定候选标签列表的候选标签;

Y3、将任务数据和候选标签分别输入到智慧客服大模型的大模型基座中,得到生成信息句向量和候选标签句向量;通过匹配算法计算生成信息句向量分别与各个候选标签句向量的相似度;将相似度最高的候选标签句向量对应的候选标签作为预测标签;

Y4、根据对话文本二的预测标签和实际标签,进行损失函数计算;当损失函数的值收敛或迭代次数达到设定次数时,下游微调阶段结束,否则调整智慧客服大模型参数,并返回分步骤Y3继续训练。

2. 根据权利要求1所述的面向客服场景的生成匹配式大模型构建方法,其特征在于:所

述步骤X2,将对话拼接文本进行自动标注处理,得到对话文本一对应的摘要,是指:包括如下分步骤:

- X21、初始化摘要集合;
- X22、将对话拼接文本拆分为n个句子;
- X23、将每个句子和摘要集合中的所有句子拼接,得到摘要拼接句;
- X24、分别计算每个摘要拼接句和对话拼接文本其余句子的最大公共子序列长度L;将L最大对应的句子从对话拼接文本中移除并存入摘要集合;
- X25、判断摘要集合中的句子数s是否达到设定句子数量:若达到,则将当前的摘要集合设定为对话文本一的摘要;否则跳至分步骤X22,继续增加摘要集合中的句子数s。

3. 根据权利要求1所述的面向客服场景的生成匹配式大模型构建方法,其特征在于:所述上下文模板包括用户信息一的内容和对话文本一的内容;用户信息一的内容包括话费余额、套餐名称、套餐资费、套餐流量。

4. 根据权利要求1所述的面向客服场景的生成匹配式大模型构建方法,其特征在于:所述预训练阶段的执行频率<领域迁移阶段的执行频率<下游微调阶段的执行频率。

5. 根据权利要求1所述的面向客服场景的生成匹配式大模型构建方法,其特征在于:所述智慧客服大模型中,大模型基座包括编码器和解码器;映射模块包括若干依次连接的全连接层。

6. 一种可读存储介质,其特征在于,其中所述存储介质存储有计算机程序,所述计算机程序当被处理器执行时使所述处理器执行权利要求1-5中任一项所述的面向客服场景的生成匹配式大模型构建方法。

7. 一种计算机设备,包括处理器以及用于存储处理器可执行程序存储器,其特征在于,所述处理器执行存储器存储的程序时,实现权利要求1-5中任一项所述的面向客服场景的生成匹配式大模型构建方法。

面向客服场景的生成匹配式大模型构建方法、介质及设备

技术领域

[0001] 本发明涉及计算机模型构建技术领域,更具体地说,涉及一种面向客服场景的生成匹配式大模型构建方法、介质及设备。

背景技术

[0002] 智慧客服大模型是在大规模知识处理基础上发展起来的一项面向客服行业的应用。在现有的客服场景中,智慧客服的核心功能是人机对话,目前主要通过自然语言生成式大模型实现。通过不同领域的语料库完成大模型的预训练后,系统将接收到的人类语言转换为序列数据并输入到大模型,经过一系列工作机制捕捉对话中的信息,还需要根据人类经验和业务规则进行一系列的判断逻辑以过滤无效信息,最终的输出转化成自然语言的形式,实现人机对话。

[0003] 但是现有智慧客服大模型构建时存在如下技术问题:

[0004] (一)常规的大模型建模方案难以适应工程化的需求:

[0005] 常规的大模型建模方案对标注数据的规模和质量有极高要求;同时,在适配不同下游任务的过程中,需要对原有模型架构进行扩展,并对原有参数进行微调,二次开发往往会带来不可忽视的额外工作量。具体到智慧客服场景中,更多复杂的实际情况仍待解决,例如,在用户来电和客服对话中,客户反映的问题通常涉及种类繁多的业务,同时还存在业务内容不断迭代更新的情况,常规的大模型预训练加微调方案很难适应具体业务的更新频率。

[0006] (二)基于深度学习的伪标签生成方案在客服场景中存在局限:

[0007] 客服对话场景涉及到各种复杂的业务问题,需要具备一定的专业知识和工作经验的客服人员才能很好地完成对话文本的标注工作,而大模型的预训练往往需要大规模的标注数据才能达到令人满意的效果。现有的基于深度学习的伪标签生成方案需要单独构建深度学习模型,并通过一定规模的标注数据完成训练后,才能达到较好的生成效果,这无疑会带来额外的开发和标注成本。对于依赖大量训练数据保证高准确率的大模型,如何充分利用人工标注信息、减小标注工作量、将人类经验与学习规则充分结合成为了急需解决的关键问题。

[0008] (三)伪标签数据包含的大量噪声会影响智慧客服大模型的训练效果:

[0009] 虽然通过深度学习模型生成伪标签数据的自动标注方式已经成为现阶段的一种主流数据增强手段,但相比于人工标注,前者很难保证标注数据的质量,往往存在一定的噪声。因此真实场景中的数据集可以分为少量人工标注的精标数据和大量带有噪声的伪标签数据,此时将整个数据集用于大模型训练,实际上是以带噪数据为主要学习对象,最小化整体损失函数的方式未必能得到预期效果。如何避免模型从大量带噪数据中学习错误的知识仍然是需要重点关注的问题。

[0010] (四)现有的智慧客服大模型功能单一,难以满足丰富的业务需求:

[0011] 现有的智慧客服大模型方案中,其核心功能局限于人机对话,并没有涉及到归因

分析等真正能够辅助客服工作人员处理来电咨询的功能。模型的技术架构通常采用生成模型结合各种规则针对用户的提问进行回复,最终输出仍然是对话文本的形式,输出结果和实际业务需求存在隔阂。

发明内容

[0012] 为克服现有技术中的缺点与不足,本发明的目的在于提供一种面向客服场景的生成匹配式大模型构建方法、介质及设备;该方法分阶段逐步地实现和优化大模型的功能,使其具备深度挖掘大规模客服文本数据知识的能力,同时对新增的业务需求和变化的业务内容具备精准迁移和快速扩展的能力;不仅能够适应不同的算力资源和数据规模,强化管理者对算法研发流程的管控,同时能够为不断更新的业务需求提供具有针对性的解决方案。

[0013] 为了达到上述目的,本发明通过下述技术方案予以实现:一种面向客服场景的生成匹配式大模型构建方法,包括依次执行的模型设定阶段、预训练阶段、领域迁移阶段和下游微调阶段;

[0014] 所述模型设定阶段,是指:设定基于Transformer架构的智慧客服大模型;所述智慧客服大模型包括大模型基座和连接在大模型基座输出端的映射模块;

[0015] 所述预训练阶段,是指:采用跨领域中文语料库的文本作为样本,对智慧客服大模型的大模型基座进行预训练,以使智慧客服大模型具备强泛化能力;

[0016] 所述领域迁移阶段,是指:采用客服场景数据作为样本;利用客服场景数据中的精标数据和通过自动标注得到的伪标签数据,对智慧客服大模型的大模型基座进行弱监督训练,以使智慧客服大模型具备强领域知识;

[0017] 所述下游微调阶段,是指:采用客服场景人工标注数据作为样本;样本根据业务需求拼接对应的业务需求提示模板,得到任务数据;利用任务数据对智慧客服大模型进行训练,以学习新业务的相关知识。

[0018] 优选地,所述领域迁移阶段,包括如下步骤:

[0019] X1、采集客服场景数据;将客服场景数据分类为用户信息一、产品信息、客服对话;客服对话包括语音转录的对话文本一;

[0020] 对话文本一包括有标签的对话文本一和无标签的对话文本一;有标签的对话文本一数量>无标签的对话文本一数量;

[0021] X2、当对话文本一为有标签的对话文本一时,以“对话文本-标签”形式形成精标数据;

[0022] 当对话文本一为无标签的对话文本一时,获取对话文本一对应的用户信息一;根据上下文模板,将用户信息一转换成文字描述,并与对话文本一进行拼接,得到对话拼接文本;将对话拼接文本进行自动标注处理,得到对话文本一对应的摘要;将对话文本一和摘要组合,以“对话文本-摘要”形式形成伪标签数据,以实现数据增强;

[0023] X3、将标签数据和伪标签数据分别输入智慧客服大模型的大模型基座;在一次前向传播过程中,计算得到标签数据和伪标签数据的损失梯度,并比较损失梯度的相似程度;

[0024] 若损失梯度方向一致,则判定为训练有效,之后进行后向传播,以完成弱监督训练;

[0025] 否则将伪标签数据的损失梯度置零,之后进行后向传播,以完成弱监督训练。

[0026] 优选地,所述步骤X2,将对话拼接文本进行自动标注处理,得到对话文本一对应的摘要,是指:包括如下分步骤:

[0027] X21、初始化摘要集合;

[0028] X22、将对话拼接文本拆分为n个句子;

[0029] X23、将每个句子和摘要集合中的所有句子拼接,得到摘要拼接句;

[0030] X24、分别计算每个摘要拼接句和对话拼接文本其余句子的最大公共子序列长度L;将L最大对应的句子从对话拼接文本中移除并存入摘要集合;

[0031] X25、判断摘要集合中的句子数s是否达到设定句子数量:若达到,则将当前的摘要集合设定为对话文本一的摘要;否则跳至分步骤X22,继续增加摘要集合中的句子数s。

[0032] 优选地,所述上下文模板包括用户信息一的内容和对话文本一的内容;用户信息一的内容包括话费余额、套餐名称、套餐资费、套餐流量。

[0033] 优选地,所述下游微调阶段,包括如下分步骤:

[0034] Y1、采集客服场景人工标注数据;人工标注数据包括有标签的对话文本二;获取对话文本二对应的用户信息二;根据业务需求设定业务需求提示模板;根据所述上下文模板,将用户信息二转换成文字描述,并与对话文本二进行拼接,之后与业务需求提示模板拼接,得到任务数据;

[0035] Y2、设定候选标签列表的候选标签;

[0036] Y3、将任务数据和候选标签分别输入到智慧客服大模型的大模型基座中,得到生成信息句向量和候选标签句向量;通过匹配算法计算生成信息句向量分别与各个候选标签句向量的相似度;将相似度最高的候选标签句向量对应的候选标签作为预测标签;

[0037] Y4、根据对话文本二的预测标签和实际标签,进行损失函数计算;当损失函数的值收敛或迭代次数达到设定次数时,下游微调阶段结束,否则调整智慧客服大模型参数,并返回分步骤Y3继续训练。

[0038] 优选地,所述预训练阶段的执行频率<领域迁移阶段的执行频率<下游微调阶段的执行频率。

[0039] 优选地,所述智慧客服大模型中,大模型基座包括编码器和解码器;映射模块包括若干依次连接的全连接层。

[0040] 一种可读存储介质,其中所述存储介质存储有计算机程序,所述计算机程序当被处理器执行时使所述处理器执行所述的面向客服场景的生成匹配式大模型构建方法。

[0041] 一种计算机设备,包括处理器以及用于存储处理器可执行程序存储器,所述处理器执行存储器存储的程序时,实现所述的面向客服场景的生成匹配式大模型构建方法。

[0042] 与现有技术相比,本发明具有如下优点与有益效果:

[0043] (一)根据现实的客服工作内容和客服数据特点将智慧客服大模型的建模全流程拆解为模型设定、预训练、领域迁移、下游微调四个阶段,分阶段逐步地实现和优化大模型的功能,使其具备深度挖掘大规模客服文本数据知识的能力,同时对新增的业务需求和变化的业务内容具备精准迁移和快速扩展的能力;多阶段建模方法中,预训练阶段通过对开放域公共语料的学习,有效解决了客服对话的冷启动问题;领域迁移阶段通过少量标注数据挖掘大量无标注对话文本的有效信息;下游微调阶段采用匹配算法快速且准确地将大模型的能力迁移到具体业务,避免了二次开发的资源消耗。分阶段的建模方法精准对标了客

服工作的整个流程,不仅能够适应不同的算力资源和数据规模,强化管理者对算法研发流程的管控,同时能够为不断更新的业务需求提供具有针对性的解决方案。

[0044] (二)采用面向对话文本的自动标注算法,从每段对话中提取摘要信息,构建“原文-摘要”形式的语料库,利用海量的原始对话数据生成大模型领域迁移训练所需要的伪标签数据;相比于独立构建深度学习模型生成伪标签数据的方案,自动标注算法更关注于数据本身的特性,通过对真实数据的分布特点进行统计分析挖掘出特定场景下的语言模式,无需额外的标注数据,极大地减少了开发工作量和算力消耗,能够充分利用现有的海量原始对话数据在短时间内生成大量对话文本的伪标签数据,极大降低了标注工作的时间成本和人力成本。

[0045] (三)采用了基于数据增强的弱监督预学习方法,在自动标注算法生成的大量伪标注数据的基础上,结合少量人工标注数据,通过弱监督的方式实现客服场景下的数据增强;人工标注的对话文本包含客服场景独有的语言模式,可视为强监督数据;使用自动标注算法能够以低成本获取到大规模的伪标签数据,但在质量上要低于人工标注数据,可视为包含噪声的弱监督数据。本发明提出的弱监督训练方法在训练数据集质量参差不齐的情况下,以数量更多的带噪数据作为模型的主要学习对象,能够基于少量的强监督数据从大量的带噪数据中挖掘有效信息,用于提升模型效果。

[0046] (四)采用基于生成式大模型的匹配算法,利用特定业务场景的标注数据完成智慧客服大模型的微调,在此基础上将业务数据和对话文本同时输入大模型进行编码,最后通过特征匹配的方式实现具体业务场景需要的功能。匹配算法,通过设置不同的提示模板训练和引导大模型在不同应用场景下的生成能力,并利用特征匹配的方式衔接大模型生成能力和业务需求,对于不断迭代的业务内容和新增的业务场景具有良好的适应能力,避免二次开发的成本消耗。

附图说明

[0047] 图1是本发明面向客服场景的生成匹配式大模型构建方法的流程图;

[0048] 图2是本发明面向客服场景的生成匹配式大模型构建方法的预训练阶段示意图;

[0049] 图3是本发明面向客服场景的生成匹配式大模型构建方法的自动标注算法流程图;

[0050] 图4是本发明面向客服场景的生成匹配式大模型构建方法的弱监督训练流程图;

[0051] 图5是本发明面向客服场景的生成匹配式大模型构建方法的下游微调阶段编码器和解码器示意图;

[0052] 图6是本发明面向客服场景的生成匹配式大模型构建方法的句向量映射示意图。

具体实施方式

[0053] 下面结合附图与具体实施方式对本发明作进一步详细的描述。

[0054] 实施例一

[0055] 本实施例一种面向客服场景的生成匹配式大模型构建方法,包括依次执行的模型设定阶段、预训练阶段、领域迁移阶段和下游微调阶段,如图1所示。

[0056] 现有大模型的应用范式是在预训练的基础上进行微调,但现实生活的应用场景

中,人类对话具有复杂且多歧义的特点,简单地套用范式无法得到满意的效果,因此本发明通过分析客服场景的工作内容和数据特性,总结出一套适用于智慧客服场景的建模方法。

[0057] 所述模型设定阶段,是指:设定基于Transformer架构的智慧客服大模型;智慧客服大模型包括大模型基座和包括若干全连接层的映射模块;大模型基座包括编码器和解码器。

[0058] 所述预训练阶段,是指:采用跨领域中文语料库的文本作为样本,对智慧客服大模型的大模型基座进行预训练,以使智慧客服大模型具备强泛化能力。预训练阶段是领域无关和任务无关的,利用开放域的中文语料库完成大模型基座的训练。

[0059] 预训练阶段,编码器是将输入的文本转化成序列数据,从中提取出语法、词性、实体等高阶语义信息并映射成特征向量;解码器将编码器输出的特征向量重新映射成数据序列,并通过数据标签提供的监督信号不断优化模型参数。在海量的中文语料库基础上,通过屏蔽词还原、主题分类、摘要提取等丰富的预训练任务学习中文语言的数据分布特性,构建中文语义空间,因此需要从多方数据源进行数据的收集和清洗,构建跨领域的中文语料库,并按照不同任务模板转化数据。经过预训练阶段输出的智慧客服大模型具备强大的中文语义理解能力和泛化能力。预训练阶段的训练方法可采用现有技术。

[0060] 所述领域迁移阶段,是指:采用客服场景数据作为样本;利用客服场景数据中的精标数据和通过自动标注得到的伪标签数据,对智慧客服大模型的大模型基座进行弱监督训练,以使智慧客服大模型具备强领域知识。

[0061] 所述下游微调阶段,是指:采用客服场景人工标注数据作为样本;样本根据业务需求拼接对应的业务需求提示模板,得到任务数据;利用任务数据对智慧客服大模型进行训练,以学习新业务的相关知识。

[0062] 具体地说,预训练阶段中,所述编码器将输入的文本转化为序列数据,从序列数据中提取出语法、词性、实体等高阶语义信息并映射成特征向量;所述解码器将编码器输出的特征向量重新映射成数据序列,并通过数据标签提供的监督信号不断优化模型参数,如图2所示。

[0063] 领域迁移阶段的核心是利用客服场景的数据实现大模型在垂直领域的知识迁移。

[0064] 所述领域迁移阶段,包括如下步骤:

[0065] X1、采集客服场景数据;将客服场景数据分类为用户信息一、产品信息、客服对话。

[0066] 其中,用户信息一的内容包括话费余额、套餐名称、套餐资费、套餐流量等信息,以及近期的自助操作、来电咨询、投诉的操作记录;产品信息包括各类套餐和服务的详细内容,例如校园套餐及其费用、免费通话时长、定向流量等套餐内容;客服对话包括语音转录的对话文本一,以及定期的人工抽检记录(一段对话文本一对应一句人工总结语句)。

[0067] 用户信息一和产品信息能够为会话的意图识别提供关键信息,客服对话则反映出从用户信息一和产品信息分析并解决问题的流程,是智慧客服大模型最主要的学习对象。

[0068] 对话文本一包括有标签的对话文本一和无标签的对话文本一;有标签的对话文本一数量>无标签的对话文本一数量;

[0069] 将有标签的对话文本一作为精标数据;对无标签的对话文本一进行自动标注,使无标签的对话文本一转化成伪标签数据。

[0070] X2、当对话文本一为有标签的对话文本一时,以“对话文本-标签”形式形成精标数

据；

[0071] 当对话文本一为无标签的对话文本一时,为了丰富当前对话的上下文环境,每个对话文本一需要和用户信息一进行匹配;本发明获取对话文本一对应的用户信息一;根据上下文模板,将用户信息一转换成文字描述,并与对话文本一进行拼接,得到对话拼接文本。所述上下文模板包括用户信息一的内容和对话文本一的内容;用户信息一的内容包括话费余额、套餐名称、套餐资费、套餐流量。假设上下文模板如表1所示,

[0072] 表1上下文模板

[0073]	字段	内容
	话费余额	20元
	套餐名称	校园学霸套餐
	套餐资费	29元/月
	套餐流量	国内通用流量30G,定向流量50G
	会话内容	你好,我的套餐流量不够用,有没有更大的流量套餐,最好是…

[0074] 对话拼接文本为:当前用户的话费余额为(20元),使用套餐为(校园学霸套餐),套餐资费为(29元每月),包含流量(国内通用流量30G,定向流量50G),当前对话内容为(你好,我的套餐流量不够用,有没有更大的流量套餐,最好是…)。

[0075] 将对话拼接文本进行自动标注处理,得到对话文本一对应的摘要;将对话文本一和摘要组合,以“对话文本-摘要”形式形成伪标签数据,以实现数据增强。具体地说,如图3所示,包括如下分步骤:

[0076] X21、初始化摘要集合;

[0077] X22、将对话拼接文本拆分为n个句子;

[0078] X23、将每个句子和摘要集合中的所有句子拼接,得到摘要拼接句;

[0079] X24、分别计算每个摘要拼接句和对话拼接文本其余句子的最大公共子序列长度L;将L最大对应的句子从对话拼接文本中移除并存入摘要集合;

[0080] X25、判断摘要集合中的句子数s是否达到设定句子数量:若达到,则将当前的摘要集合设定为对话文本一的摘要;否则跳至分步骤X22,继续增加摘要集合中的句子数s。

[0081] X3、将标签数据和伪标签数据分别输入智慧客服大模型的大模型基座,以进行弱监督训练;弱监督训练是指:如图4所示,在一次前向传播过程中,计算得到标签数据和伪标签数据的损失梯度,并比较损失梯度的相似程度:

[0082] 若损失梯度方向一致,则判定为训练数据的优化方向与干净数据的优化方向一致,训练有效,之后进行后向传播,以完成弱监督训练;

[0083] 否则判定为噪声影响了干净数据训练,将伪标签数据的损失梯度置零,之后进行后向传播,以完成弱监督训练。

[0084] 下游微调阶段的目标是面向不同的业务需求提供解决方案,智慧客服大模型从少量的人工标注数据中学习新业务的相关知识,提升具体任务的性能表现。

[0085] 下游微调阶段,包括如下分步骤:

[0086] Y1、采集客服场景人工标注数据;人工标注数据包括有标签的对话文本二;获取对话文本二对应的用户信息二;根据业务需求设定业务需求提示模板;根据所述上下文模板,将用户信息二转换成文字描述,并与对话文本二进行拼接,之后与业务需求提示模板拼接,

得到任务数据;

[0087] 用户对话归因分析作为智慧客服场景中的重要功能,能够为企业提供关键的分析结果,辅助客服工作人员及时发现和解决问题,最终达到提升用户满意度的目的。本实施例以用户对话归因分析为例,说明下游微调阶段的工作内容;归因分析是指根据用户的来电和会话内容,客服工作人员结合现有的信息将当前会话归类为具体的业务或产品。如果需要从某一时间段内的用户来电和会话分析出哪些产品的热度最高,可构建产品归因提示模板;如果需要定位到用户投诉和产品发生故障的原因,可构建投诉归因提示模板;对于其他的业务需求,可针对业务数据特点构建不同类型的提示模板,为智慧客服大模型的输入数据提供辅助信息。

[0088] 当业务需求为产品归因时,业务需求提示模板的内容为:“当前对话属于什么产品?”;当业务需求为投诉归因时,业务需求提示模板的内容为:“当前对话的投诉内容是什么?”

[0089] 以投诉归因为例,任务数据为:

[0090] “当前用户的话费余额为(20元),使用套餐为(校园学霸套餐),套餐资费为(29元每月),包含流量(国内通用流量30G,定向流量50G),当前对话内容为(你好,我的套餐流量不够用,有没有更大的流量套餐,最好是…)。当前对话的投诉内容是什么?”

[0091] Y2、从业务系统中根据关键字查询到相关的产品和活动作为候选标签,并保存到候选标签列表;

[0092] Y3、将任务数据输入到智慧客服大模型的编码器中,智慧客服大模型的解码器输出解码器生成信息;将解码器生成信息输入智慧客服大模型的映射模块,经过映射模块多层全连接层的映射,得到生成信息句向量,如图5所示;将候选标签输入智慧客服大模型,得到候选标签句向量;

[0093] 编码器基于注意力机制对输入文本的词性、语法、语境、上下文关联以及相对位置等关键信息进行深度挖掘,并编码成隐式特征向量,作为解码器的输入。解码器负责将隐式特征向量重构成输入文本的标签。通过编码-解码的工作流程从大规模“文本-标签”样本对中挖掘出客服场景下的语言模式和语境常识。

[0094] 句向量映射的目的是把任务数据和候选标签文本同时输入大模型编码并进行相似度匹配,实现任务数据和标签的高效匹配,该部分需要在解码器之后构建包含多个全连接层的映射模块,将解码器的生成语句映射为固定长度的句向量。

[0095] 通过匹配算法计算生成信息句向量分别与各个候选标签句向量的余弦相似度;将相似度最高的候选标签句向量对应的候选标签作为预测标签,如图6所示。

[0096] Y4、根据对话文本二的预测标签和实际标签,进行损失函数计算;当损失函数的值收敛或迭代次数达到设定次数时,下游微调阶段结束,否则调整智慧客服大模型参数,并返回分步骤Y3继续训练。

[0097] 实施例二

[0098] 本实施例一种可读存储介质,其中所述可读存储介质存储有计算机程序,所述计算机程序当被处理器执行时使所述处理器执行实施例一所述的面向客服场景的生成匹配式大模型构建方法。

[0099] 实施例三

[0100] 本实施例一种计算机设备,包括处理器以及用于存储处理器可执行程序存储器,所述处理器执行存储器存储的程序时,实现实施例一所述的面向客服场景的生成匹配式大模型构建方法。

[0101] 上述实施例为本发明较佳的实施方式,但本发明的实施方式并不受上述实施例的限制,其他的任何未背离本发明的精神实质与原理下所作的改变、修饰、替代、组合、简化,均应为等效的置换方式,都包含在本发明的保护范围之内。

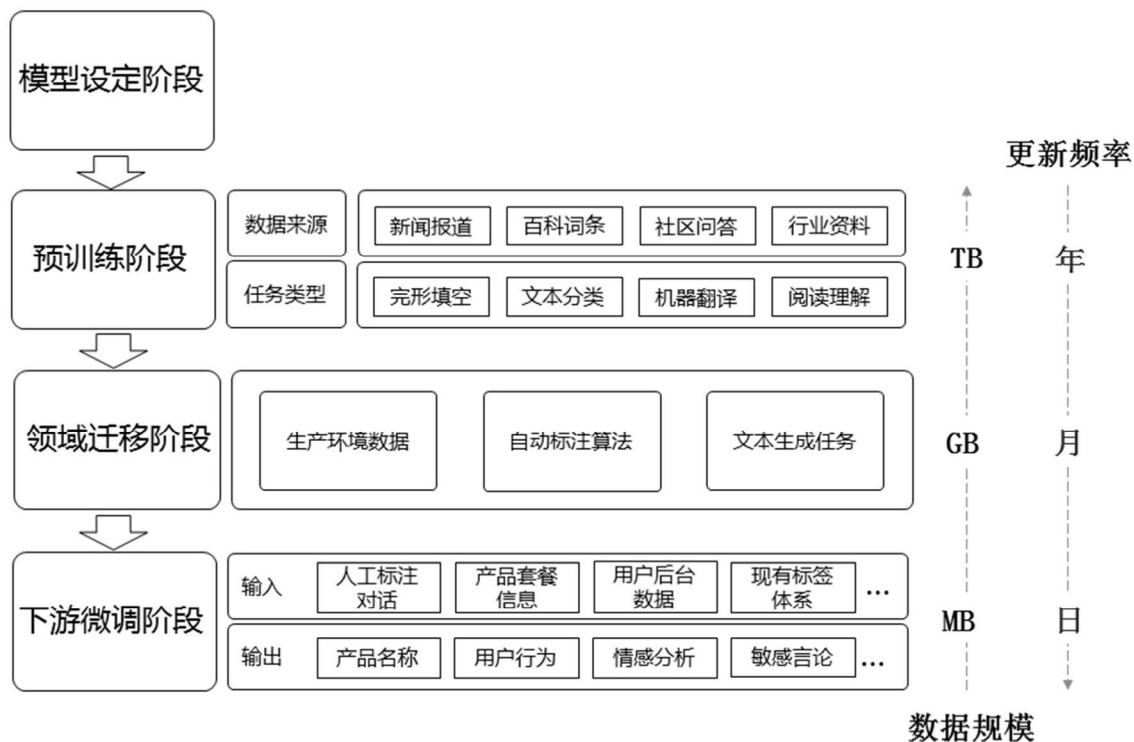


图1

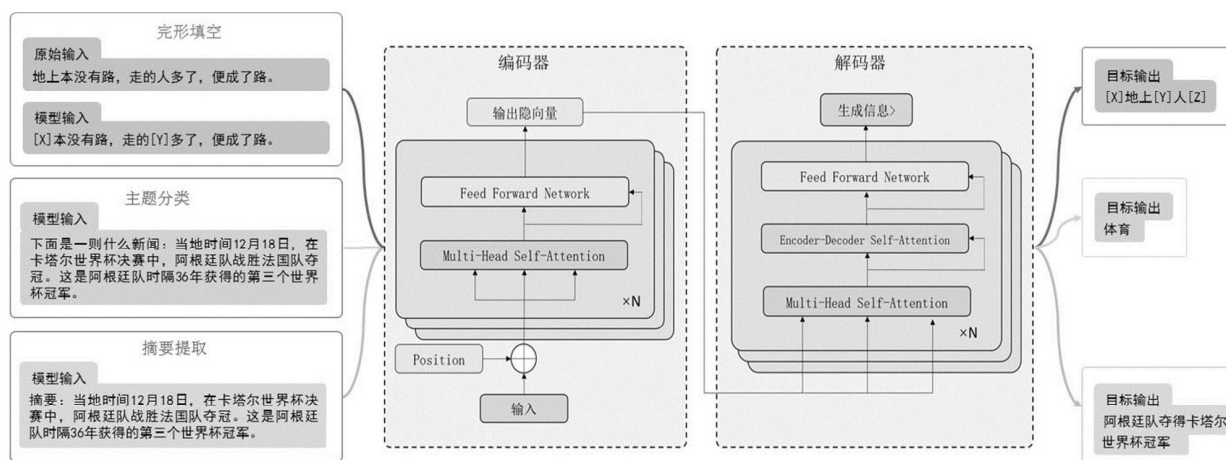


图2

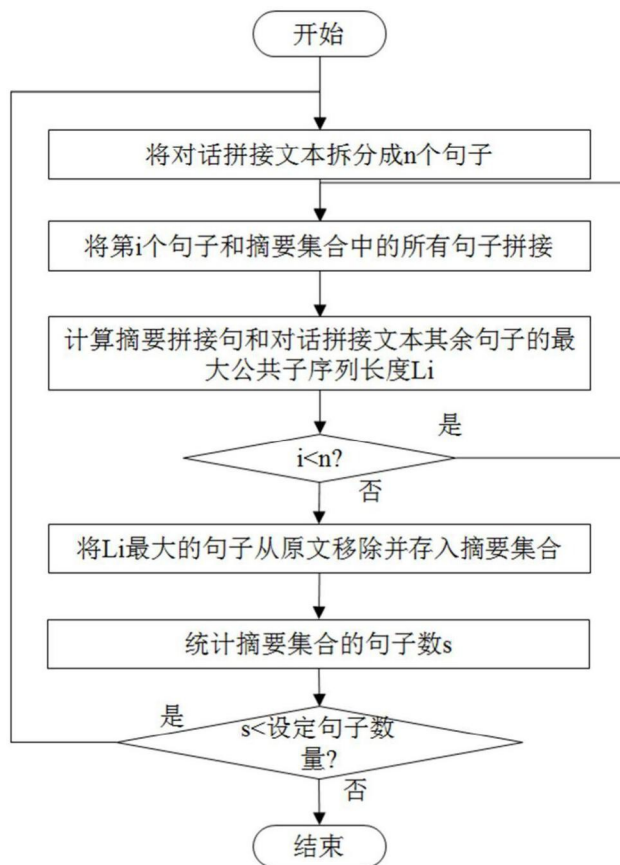


图3

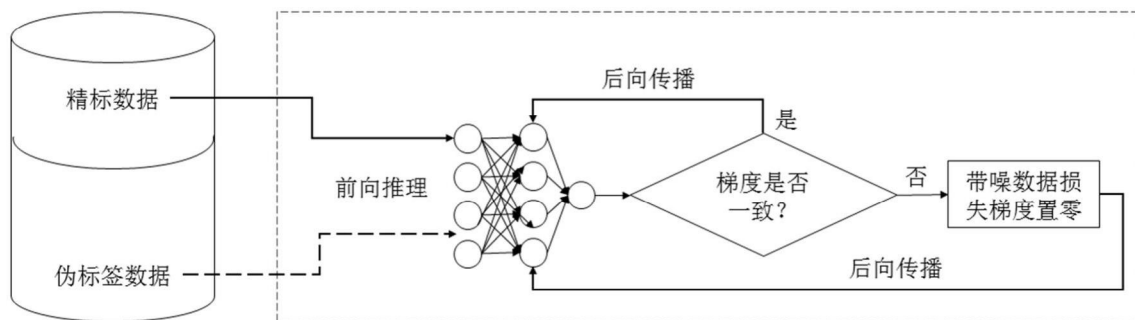


图4

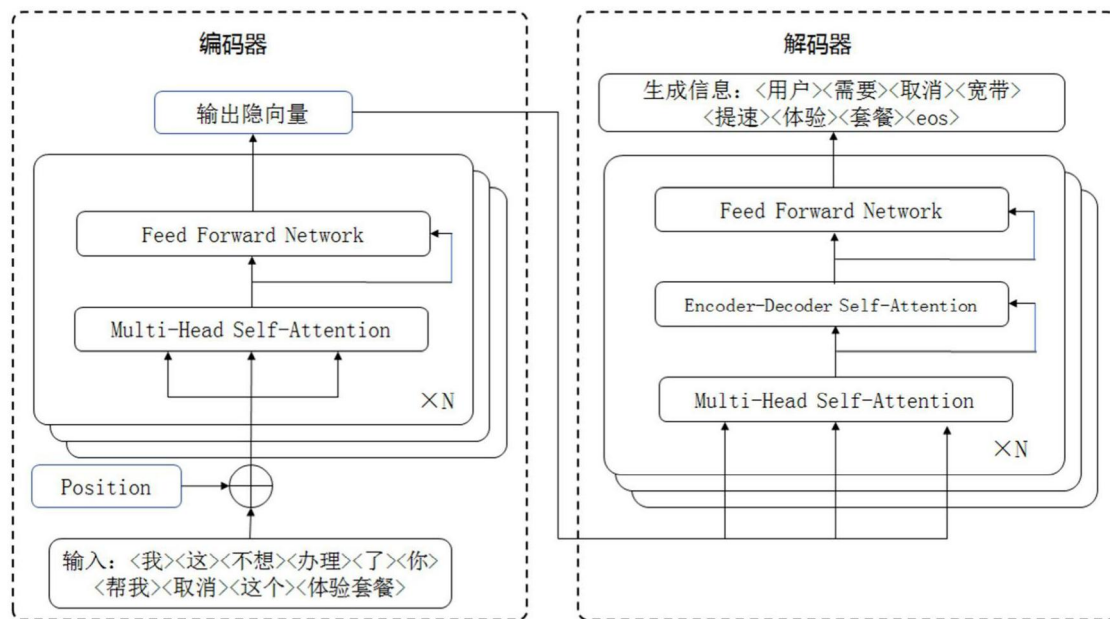


图5

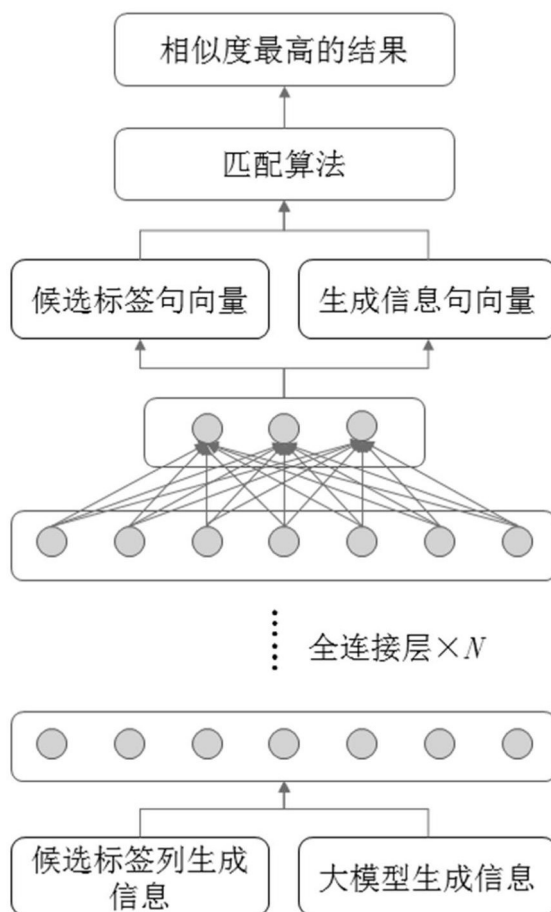


图6